

CS 229, Spring 2020

Problem Set #2

Due Wednesday, May 13 at 11:59 pm on Gradescope.

Notes: (1) These questions require thought, but do not require long answers. Please be as concise as possible. (2) If you have a question about this homework, we encourage you to post your question on our Piazza forum, at <https://piazza.com/stanford/spring2020/cs229>. (3) This quarter, Spring 2020, students may submit in pairs. If you missed the first lecture or are unfamiliar with the collaboration or honor code policy, please read the policy on the course website before starting work. (4) For the coding problems, you may not use any libraries except those defined in the provided `environment.yml` file. In particular, ML-specific libraries such as `scikit-learn` are not permitted. (5) To account for late days, the due date is Wednesday, May 13 at 11:59 pm. If you submit after Wednesday, May 13 at 11:59 pm, you will begin consuming your late days. If you wish to submit on time, submit before Wednesday, May 13 at 11:59 pm.

All students must submit an electronic PDF version of the written questions. We highly recommend typesetting your solutions via $\text{\LaTeX}$. All students must also submit a zip file of their source code to Gradescope, which should be created using the `make.zip.py` script. You should make sure to (1) restrict yourself to only using libraries included in the `environment.yml` file, and (2) make sure your code runs without errors. Your submission may be evaluated by the auto-grader using a private test set, or used for verifying the outputs reported in the writeup.

1. [15 points] Logistic Regression: Training stability

In this problem, we will be delving deeper into the workings of logistic regression. The goal of this problem is to help you develop your skills debugging machine learning algorithms (which can be very different from debugging software in general).

We have provided an implementation of logistic regression in `src/stability/stability.py`, and two labeled datasets A and B in `src/stability/ds1.a.csv` and `src/stability/ds1.b.csv`.

Please do not modify the code for the logistic regression training algorithm for this problem. First, run the given logistic regression code to train two different models on A and B . You can run the code by simply executing `python stability.py` in the `src/stability` directory.

- [2 points] What is the most notable difference in training the logistic regression model on datasets A and B ?
- [5 points] Investigate why the training procedure behaves unexpectedly on dataset B , but not on A . Provide hard evidence (in the form of math, code, plots, etc.) to corroborate your hypothesis for the misbehavior. Remember, you should address why your explanation does *not* apply to A .

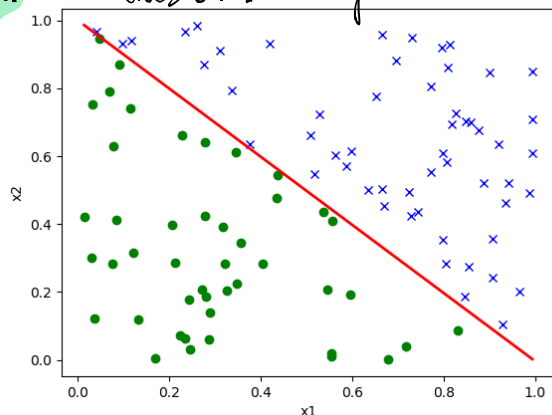
Hint: The issue is not a numerical rounding or over/underflow error.

- [5 points] For each of these possible modifications, state whether or not it would lead to the provided training algorithm converging on datasets such as B . Justify your answers.
 - Using a different constant learning rate.
 - Decreasing the learning rate over time (e.g. scaling the initial learning rate by $1/t^2$, where t is the number of gradient descent iterations thus far).
 - Linear scaling of the input features.
 - Adding a regularization term $\|\theta\|_2^2$ to the loss function.
 - Adding zero-mean Gaussian noise to the training data or labels.
- [3 points] Are support vector machines vulnerable to datasets like B ? Why or why not? Give an informal justification.

(a) Data set B takes much larger number of iterations as well as time to train. It won't converge at least at 60,000,000 iterations.

(b) See the decision boundary of data set B after 302,000 iterations. (fig 1, a)

a. decision boundary.



b. loss curve

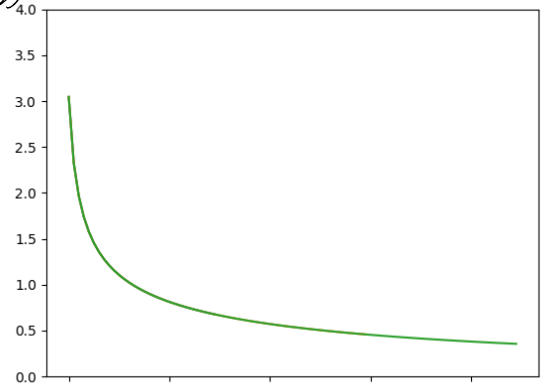


figure 1

$\times 10,000$ iterations

Actually, it's already reasonably well. The reason of non-convergence maybe that the learning rate is too big for the convergence criteria. Let's check by taking a look at loss curve (fig 1. b). From b, we can see the curve is quite smooth and the loss is still descending. Look back to fig 1. a, it seems like a perfect separation case where theta would constantly increase and loss will converge on 0 after infinite iterations. Check this by plotting l_2 norm of theta. (fig 2. a) and test the accuracy of model

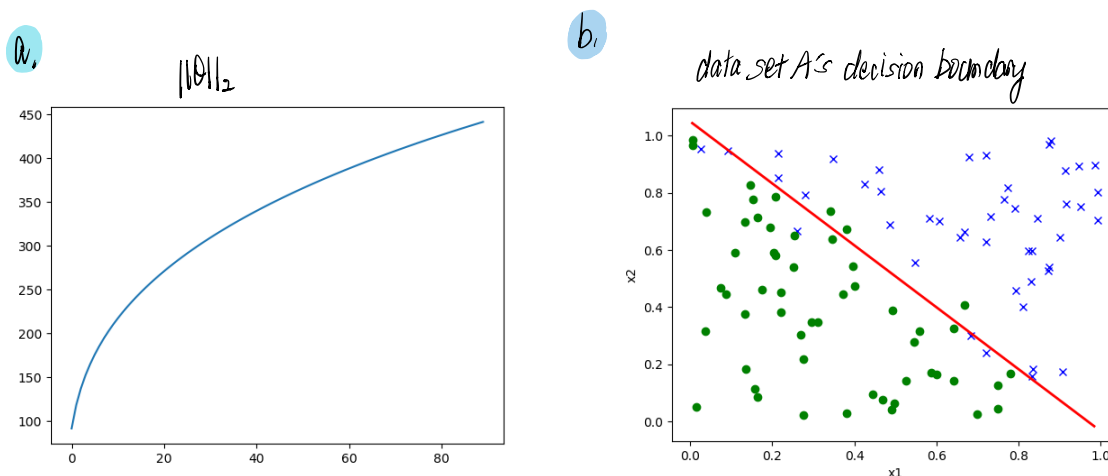


figure 2

The accuracy is 1 since the first 10,000 iterations, and theta is constantly increasing. Data set B has a problem of perfect separation. Let double check if data set A has. (fig 2. b)

(c) i. **No**. No matter setting a larger or smaller learning rate, the model won't converge until $\|\theta\|_2 \rightarrow \infty$. It won't solve this problem.

ii. **No**. The possibility of learning rate being too large has been ruled out in the analysis where the loss curve is constantly decreasing, instead of fluctuating.

iii. **No**. After linear scaling, X will still be perfectly separable. It won't solve.

iv. **Yes**. It constraints the weights θ from going too large. So the iteration will stop somewhere when θ becomes large.

v. Maybe. Gaussian noise would work if it makes X to be linearly unseparable. But it won't work if it doesn't.

(d) No. SVM works better when cases are perfectly separable. In general, SVM maximize the margin around the separating hyperplane. When X is linearly separable w, b are obtained by

$$\max_{w, b} \min_{i=1} \frac{y_i(w^T x_i + b)}{\|w\|}$$

where $w^T x_i + b$ denote the data x_i 's distance to the hyperplane.

2. [22 points] Spam classification

In this problem, we will use the naive Bayes algorithm and an SVM to build a spam classifier.

In recent years, spam on electronic media has been a growing concern. Here, we'll build a classifier to distinguish between real messages, and spam messages. For this class, we will be building a classifier to detect SMS spam messages. We will be using an SMS spam dataset developed by Tiago A. Almeida and Jos Mara Gmez Hidalgo which is publicly available on <http://www.dt.fee.unicamp.br/~tiago/smsspamcollection>¹

We have split this dataset into training and testing sets and have included them in this assignment as `src/spam/spam_train.tsv` and `src/spam/spam_test.tsv`. See `src/spam/spam_readme.txt` for more details about this dataset. Please refrain from redistributing these dataset files. The goal of this assignment is to build a classifier from scratch that can tell the difference the spam and non-spam messages using the text of the SMS message.

- (a) [5 points] Implement code for processing the the spam messages into numpy arrays that can be fed into machine learning models. Do this by completing the `get_words`, `create_dictionary`, and `transform_text` functions within our provided `src/spam.py`. Do note the corresponding comments for each function for instructions on what specific processing is required.

The provided code will then run your functions and save the resulting dictionary into `spam_dictionary` and a sample of the resulting training matrix into `spam_sample_train_matrix`.

In your writeup, report the vocabular size after the pre-processing step. You do not need to include any other output for this subquestion.

- (b) [10 points] In this question you are going to implement a naive Bayes classifier for spam classification with **multinomial event model** and Laplace smoothing.

Code your implementation by completing the `fit_naive_bayes_model` and `predict_from_naive_bayes_model` functions in `src/spam/spam.py`.

Now `src/spam/spam.py` should be able to train a Naive Bayes model, compute your prediction accuracy and then save your resulting predictions to `spam_naive_bayes_predictions`.

In your writeup, report the accuracy of the trained model on the **test set**.

Remark. If you implement naive Bayes the straightforward way, you will find that the computed $p(x|y) = \prod_i p(x_i|y)$ often equals zero. This is because $p(x|y)$, which is the product of many numbers less than one, is a very small number. The standard computer representation of real numbers cannot handle numbers that are too small, and instead rounds them off to zero. (This is called “underflow.”) You'll have to find a way to compute Naive Bayes' predicted class labels without explicitly representing very small numbers such as $p(x|y)$. [**Hint:** Think about using logarithms.]

- (c) [5 points] Intuitively, some tokens may be particularly indicative of an SMS being in a particular class. We can try to get an informal sense of how indicative token i is for the SPAM class by looking at:

$$\log \frac{p(x_j = i \mid y = 1)}{p(x_j = i \mid y = 0)} = \log \left(\frac{P(\text{token } i \mid \text{email is SPAM})}{P(\text{token } i \mid \text{email is NOTSPAM})} \right).$$

Complete the `get_top_five_naive_bayes_words` function within the provided code using the above formula in order to obtain the 5 most indicative tokens.

¹Almeida, T.A., Gmez Hidalgo, J.M., Yamakami, A. Contributions to the Study of SMS Spam Filtering: New Collection and Results. Proceedings of the 2011 ACM Symposium on Document Engineering (DOCENG'11), Mountain View, CA, USA, 2011.

Report the top five words in your writeup.

- (d) [2 points] Support vector machines (SVMs) are an alternative machine learning model that we discussed in class. We have provided you an SVM implementation (using a radial basis function (RBF) kernel) within `src/spam/svm.py` (You should not need to modify that code).

One important part of training an SVM parameterized by an RBF kernel (a.k.a Gaussian kernel) is choosing an appropriate kernel radius parameter.

Complete the `compute_best_svm_radius` by writing code to compute the best SVM radius which maximizes accuracy on the validation dataset. Report the best kernel radius you obtained in the writeup.

(a) dictionary size: 1721

(b) Accuracy = 0.734 on testing set.

(c) claim, won, prize, urgent, awarded

(d) The best radius was 0.1, accuracy is 0.97.

3. [18 points] Constructing kernels

In class, we saw that by choosing a kernel $K(x, z) = \phi(x)^T \phi(z)$, we can implicitly map data to a high dimensional space, and have a learning algorithm (e.g SVM or logistic regression) work in that space. One way to generate kernels is to explicitly define the mapping ϕ to a higher dimensional space, and then work out the corresponding K .

However in this question we are interested in direct construction of kernels. I.e., suppose we have a function $K(x, z)$ that we think gives an appropriate similarity measure for our learning problem, and we are considering plugging K into the SVM as the kernel function. However for $K(x, z)$ to be a valid kernel, it must correspond to an inner product in some higher dimensional space resulting from some feature mapping ϕ . Mercer's theorem tells us that $K(x, z)$ is a (Mercer) kernel if and only if for any finite set $\{x^{(1)}, \dots, x^{(n)}\}$, the square matrix $K \in \mathbb{R}^{n \times n}$ whose entries are given by $K_{ij} = K(x^{(i)}, x^{(j)})$ is symmetric and positive semidefinite. You can find more details about Mercer's theorem in the notes, though the description above is sufficient for this problem. In this question we are interested to see which operations preserve the validity of kernels.

Let K_1, K_2 be kernels over $\mathbb{R}^d \times \mathbb{R}^d$, let $a \in \mathbb{R}^+$ be a positive real number, let $f: \mathbb{R}^d \mapsto \mathbb{R}$ be a real-valued function, let $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^p$ be a function mapping from $\mathbb{R}^d$ to $\mathbb{R}^p$, let K_3 be a kernel over $\mathbb{R}^p \times \mathbb{R}^p$, and let $p(x)$ a polynomial over x with *positive* coefficients.

For each of the functions K below, state whether it is necessarily a kernel. If you think it is, prove it; if you think it isn't, give a counter-example.

- (a) [1 points] $K(x, z) = K_1(x, z) + K_2(x, z)$
- (b) [1 points] $K(x, z) = K_1(x, z) - K_2(x, z)$
- (c) [1 points] $K(x, z) = aK_1(x, z)$
- (d) [1 points] $K(x, z) = -aK_1(x, z)$
- (e) [5 points] $K(x, z) = K_1(x, z)K_2(x, z)$
- (f) [3 points] $K(x, z) = f(x)f(z)$
- (g) [3 points] $K(x, z) = K_3(\phi(x), \phi(z))$
- (h) [3 points] $K(x, z) = p(K_1(x, z))$

[Hint: For part (e), the answer is that K is indeed a kernel. You still have to prove it, though. (This one may be harder than the rest.) This result may also be useful for another part of the problem.]

Let $u \in \mathbb{R}^d$ to be an arbitrary vector. The quadratic form of $u^T K u = a_i \in \mathbb{R}$.

(a) Yes. If K_1 and K_2 are kernel. then they are PSD

$$u^T K u = u^T K_1 u + u^T K_2 u = a_1 + a_2 \geq 0$$

(b) No. If the quadratic form of K_2 is larger than K_1 's then $a_1 - a_2 < 0$.

$$u^T K u = u^T K_1 u - u^T K_2 u = a_1 - a_2$$

(c) Yes. Since $a > 0$,

$$u^T K u = a \cdot (u^T K_1 u) = a \cdot a_1 \geq 0.$$

(d) No. A PSD multiplied by a negative number is not PSD

$$w^T K u = (-a)(u^T K u) = -a a_1 \leq 0,$$

(e) Note that K is a kernel $\Leftrightarrow K \succeq 0 \Leftrightarrow \overset{\text{kernel function}}{K(x, z)} = \phi(x)^T \phi(z)$

$$\begin{aligned} K_1(x, z) K_2(x, z) &= \phi_1^T(x) \phi_1(z) \phi_2^T(x) \phi_2(z) \\ &= \left(\sum_{i=1}^d \phi_1^{(i)}(x) \phi_1^{(i)}(z) \right) \left(\sum_{j=1}^d \phi_2^{(j)}(x) \phi_2^{(j)}(z) \right) \\ &= \sum_{i=1}^d \sum_{j=1}^d \phi_1^{(i)}(x) \cdot \phi_2^{(j)}(x) \cdot \phi_1^{(i)}(z) \cdot \phi_2^{(j)}(z) \end{aligned}$$

If let $\phi_1^{(i)}(x) \phi_2^{(j)}(x) = \phi_3^{(i, j)}(x)$

Then $K_1(x, z) K_2(x, z) = \sum_{i, j} \phi_3^{(i, j)}(x) \cdot \phi_3^{(i, j)}(z) = \phi_3^T(x) \phi_3(z)$

where $\phi_3(x) \in \mathbb{R}^{\binom{n}{2}}$

(f) Yes

$$\begin{aligned} w^T K u &= \sum_{i=1}^d \sum_{j=1}^d u^{(i)} f(x^{(i)}) f(z^{(j)}) u^{(j)} \\ &= \sum_{k=1}^d \left[u^{(k)} f(x^{(k)}) \right]^2 + 2 \sum_{i \neq j} (u^{(i)} f(x^{(i)})) (u^{(j)} f(x^{(j)})) \\ &= \left(\sum_{k=1}^d (u^{(k)} f(x^{(k)})) \right)^2 \geq 0 \end{aligned}$$

(g) Yes

$$w^T K u = \sum_{i=1}^d \sum_{j=1}^d u^{(i)} K_3(\phi(x^{(i)}), \phi(z^{(j)})) u^{(j)}$$

Let $x' = \phi(x^{(i)}) \in \mathbb{R}^p$, then $K_3(\phi(x^{(i)}), \phi(z^{(j)})) = K_3(x', z')$

Since K_3 is a kernel, then it's a kernel

(h) **Yes**. Let $p(x) = \sum_{k=0}^n a_k x^k$, where $a_k > 0$

$$\text{Then } K(x, z) = p(K_1(x, z)) = \sum_{k=0}^n a_k (K_1(x, z))^k$$

$$u^T K(x, z) u = \sum_{k=0}^n a_k (u^T K_1(x, z)^k u)$$

using (e), we know that $K_1(x, z)^k$ is a kernel. then $u^T K_1(x, z)^k u = m_k \geq 0$.

$$u^T K(x, z) u = \sum_{k=0}^n a_k m_k \geq 0$$

4. [15 points] Kernelizing the Perceptron

Let there be a binary classification problem with $y \in \{0, 1\}$. The perceptron uses hypotheses of the form $h_\theta(x) = g(\theta^T x)$, where $g(z) = \text{sign}(z) = 1$ if $z \geq 0$, 0 otherwise. In this problem we will consider a stochastic gradient descent-like implementation of the perceptron algorithm where each update to the parameters θ is made using only one training example. However, unlike stochastic gradient descent, the perceptron algorithm will only make one pass through the entire training set. The update rule for this version of the perceptron algorithm is given by

$$\theta^{(i+1)} := \theta^{(i)} + \alpha(y^{(i+1)} - h_{\theta^{(i)}}(x^{(i+1)}))x^{(i+1)}$$

where $\theta^{(i)}$ is the value of the parameters after the algorithm has seen the first i training examples. Prior to seeing any training examples, $\theta^{(0)}$ is initialized to $\vec{0}$.

- (a) [3 points] Let K be a kernel corresponding to some very high-dimensional feature mapping ϕ . Suppose ϕ is so high-dimensional (say, ∞ -dimensional) that it's infeasible to ever represent $\phi(x)$ explicitly. Describe how you would apply the “kernel trick” to the perceptron to make it work in the high-dimensional feature space ϕ , but without ever explicitly computing $\phi(x)$. [Note: You don't have to worry about the intercept term. If you like, think of ϕ as having the property that $\phi_0(x) = 1$ so that this is taken care of.] Your description should specify:

- i. [1 points] How you will (implicitly) represent the high-dimensional parameter vector $\theta^{(i)}$, including how the initial value $\theta^{(0)} = 0$ is represented (note that $\theta^{(i)}$ is now a vector whose dimension is the same as the feature vectors $\phi(x)$);
- ii. [1 points] How you will efficiently make a prediction on a new input $x^{(i+1)}$. I.e., how you will compute $h_{\theta^{(i)}}(x^{(i+1)}) = g(\theta^{(i)T} \phi(x^{(i+1)}))$, using your representation of $\theta^{(i)}$; and
- iii. [1 points] How you will modify the update rule given above to perform an update to θ on a new training example $(x^{(i+1)}, y^{(i+1)})$; i.e., using the update rule corresponding to the feature mapping ϕ :

$$\theta^{(i+1)} := \theta^{(i)} + \alpha(y^{(i+1)} - h_{\theta^{(i)}}(x^{(i+1)}))\phi(x^{(i+1)})$$

- (b) [10 points] Implement your approach by completing the `initial_state`, `predict`, and `update_state` methods of `src/perceptron/perceptron.py`.

We provide three functions to be used as kernel, a dot-product kernel defined as:

$$K(x, z) = x^\top z, \tag{1}$$

a radial basis function (RBF) kernel, defined as:

$$K(x, z) = \exp\left(-\frac{\|x - z\|_2^2}{2\sigma^2}\right), \tag{2}$$

and finally the following function:

$$K(x, z) = \begin{cases} -1 & x = z \\ 0 & x \neq z \end{cases} \tag{3}$$

Note that the last function is not a kernel function (since its corresponding matrix is not a PSD matrix). However, we are still interested to see what happens when

the kernel is invalid. Run `src/perceptron/perceptron.py` to train kernelized perceptrons on `src/perceptron/train.csv`. The code will then test the perceptron on `src/perceptron/test.csv` and save the resulting predictions in the `src/perceptron/` folder. Plots will also be saved in `src/perceptron/`.

Include the three plots (corresponding to each of the kernels) in your writeup, and indicate which plot belongs to which function.

- (c) [2 points] One of the choices in Q4b completely fails, one works a bit, and one works well in classifying the points. Discuss the performance of different choices and why do they fail or perform well?

(a) i Recall that kernel function $K(x, z) = \phi^T(x) \cdot \phi(z)$.

$$\text{And } \theta^{(k)} = \sum_{i=1}^k (\alpha(y^{(i)} - h_{\theta^{(k)}}(x^{(i)})) \cdot \phi(x^{(i)})) = \sum_{i=1}^k \beta_i \phi(x^{(i)})$$

$$\text{where } \beta_i := \alpha(y^{(i)} - h_{\theta^{(k)}}(x^{(i)}))$$

$$\text{ii. } h_{\theta^{(k)}}(x^{(i+1)}) = g(\theta^{(k)T} \phi(x^{(i+1)}))$$

$$\begin{aligned} &= g\left(\sum_{j=1}^i \beta_j \phi(x^{(j)})^T \cdot \phi(x^{(i+1)})\right) \\ &= g\left(\sum_{j=1}^i \beta_j K(x^{(j)}, x^{(i+1)})\right) \end{aligned}$$

Note that $K(x^{(j)}, x^{(i+1)})$ is precomputed in practice.

$$\text{iii. } \theta^{(i+1)} = \theta^{(i)} + \beta_{i+1} \phi(x^{(i+1)})$$

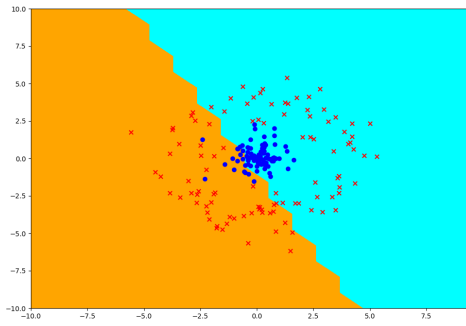
We just need to update $\beta_{i+1} = \alpha(y^{(i+1)} - h_{\theta^{(i)}}(x^{(i+1)}))$

from (ii), $h_{\theta^{(i)}}(x^{(i+1)})$ can be computed. Thus β_i can be updated.

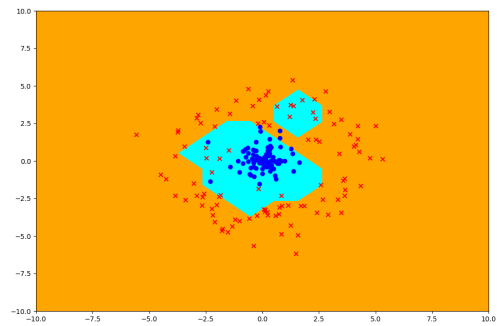
When computing $h_{\theta^{(i+1)}}(x^{(i+2)}) = g\left(\sum_{j=1}^{i+1} \beta_j K(x^{(j)}, x^{(i+2)})\right)$, we only

need to compute $\beta_1, \beta_2, \dots, \beta_{i+1}$ and $K(x^{(i)}, x^{(i+j)})$ $j=1, 2, \dots, i+1$.

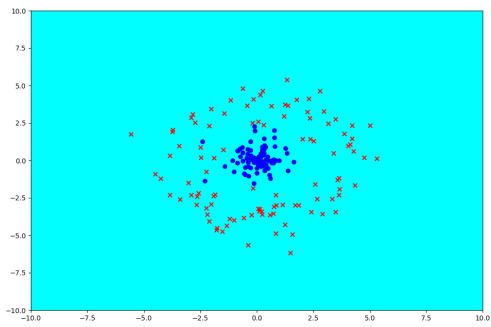
(b)



Kernel 1



Kernel 2



Kernel 3.

(C) Choice 1. works a little bit. That's because the kernel

$$K(x, z) = x^T z$$

is a kernel where $\phi(x) = x$. Which does not include new features.

The decision boundary will be linear and the flexibility of the model is not increased by this model.

Choice 3 doesn't work ($K(x, z) = \begin{cases} 1 & x=z \\ 0 & x \neq z \end{cases}$)

Consider predicting an example x .

$$h_\theta(x) = \sum_{i=1}^n \beta_i K(x^{(i)}, x)$$

if $x \neq x^{(j)}$ $j \in \{1, 2, \dots, n\}$. then $K(x^{(j)}, x) = 0$, $h_\theta(x) = 0$.

if $x = x^{(j)}$ $j \in \{1, 2, \dots, n\}$. then $h_\theta(x) = \beta_j = \sum_{i=1}^j \alpha (y^{(i)} - h_\theta(x^{(i)}))$

This kernel is making the perceptron to not be able to learn anything

It predicts 0 if the sample is not seen before, and it predicts

$\alpha \sum_{i=1}^j y^{(i)}$ if the sample was first seen in j th position.

5. [25 points] Bayesian Interpretation of Regularization

Background: In Bayesian statistics, almost every quantity is a random variable, which can either be observed or unobserved. For instance, parameters θ are generally unobserved random variables, and data x and y are observed random variables. The joint distribution of all the random variables is also called the *model* (e.g., $p(x, y, \theta)$). Every unknown quantity can be estimated by conditioning the model on all the observed quantities. Such a conditional distribution over the unobserved random variables, conditioned on the observed random variables, is called the *posterior distribution*. For instance $p(\theta|x, y)$ is the posterior distribution in the machine learning context. A consequence of this approach is that we are required to endow our model parameters, i.e., $p(\theta)$, with a *prior distribution*. The prior probabilities are to be assigned *before* we see the data—they capture our prior beliefs of what the model parameters might be before observing any evidence.

In the purest Bayesian interpretation, we are required to keep the entire posterior distribution over the parameters all the way until prediction, to come up with the *posterior predictive distribution*, and the final prediction will be the expected value of the posterior predictive distribution. However in most situations, this is computationally very expensive, and we settle for a compromise that is *less pure* (in the Bayesian sense).

The compromise is to estimate a point value of the parameters (instead of the full distribution) which is the mode of the posterior distribution. Estimating the mode of the posterior distribution is also called *maximum a posteriori estimation* (MAP). That is,

$$\theta_{\text{MAP}} = \arg \max_{\theta} p(\theta|x, y).$$

Compare this to the *maximum likelihood estimation* (MLE) we have seen previously:

$$\theta_{\text{MLE}} = \arg \max_{\theta} p(y|x, \theta).$$

In this problem, we explore the connection between MAP estimation, and common regularization techniques that are applied with MLE estimation. In particular, you will show how the choice of prior distribution over θ (e.g., Gaussian or Laplace prior) is equivalent to different kinds of regularization (e.g., L_2 , or L_1 regularization). You will also explore how regularization strengths affect generalization in part (d).

- (a) [3 points] Show that $\theta_{\text{MAP}} = \arg \max_{\theta} p(y|x, \theta)p(\theta)$ if we assume that $p(\theta) = p(\theta|x)$. The assumption that $p(\theta) = p(\theta|x)$ will be valid for models such as linear regression where the input x are not explicitly modeled by θ . (Note that this means x and θ are marginally independent, but not conditionally independent when y is given.)
- (b) [5 points] Recall that L_2 regularization penalizes the L_2 norm of the parameters while minimizing the loss (i.e., negative log likelihood in case of probabilistic models). Now we will show that MAP estimation with a zero-mean Gaussian prior over θ , specifically $\theta \sim \mathcal{N}(0, \eta^2 I)$, is equivalent to applying L_2 regularization with MLE estimation. Specifically, show that for some scalar λ ,

$$\theta_{\text{MAP}} = \arg \min_{\theta} -\log p(y|x, \theta) + \lambda \|\theta\|_2^2. \quad (4)$$

Also, what is the value of λ ?

- (c) [7 points] Now consider a specific instance, a linear regression model given by $y = \theta^T x + \epsilon$ where $\epsilon \sim \mathcal{N}(0, \sigma^2)$. Assume that the random noise $\epsilon^{(i)}$ is independent for every training

example $x^{(i)}$. Like before, assume a Gaussian prior on this model such that $\theta \sim \mathcal{N}(0, \eta^2 I)$. For notation, let X be the design matrix of all the training example inputs where each row vector is one example input, and $\vec{y}$ be the column vector of all the example outputs.

Come up with a closed form expression for θ_{MAP} .

- (d) [5 points] Coding question: double descent phenomenon and the effect of regularization.

The Bayesian perspective provides a justification of using the L_2 regularization term as in equation (4), and it also provides a formula for the optimal choice of λ , assuming that we know the true prior for θ . For real-world applications, we often do not know the true prior, so λ is tuned based on the performance on the validation dataset. In this problem, you will be asked to verify empirically that the choice of λ you derived in part (c) is close to optimal, when the generating process of the data, θ , and the label y are as exactly described in part (c).

Meanwhile, you will empirically observe an interesting phenomenon, often called double descent, which was first discovered in 1970s and recently returned to spotlight. Sample-wise double descent is the phenomenon that the validation loss of some learning algorithm or estimator does not monotonically decrease as we have more training examples, but instead has a curve with two U-shaped parts. Model-wise double descent is a similar phenomenon as we increase the size of the model. For simplicity, here we only focus on the sample-wise double descent.

You will be asked to train on various datasets by minimizing the regularized cost function with various choices of λ :

$$-\log p(\vec{y}|X, \theta) + \lambda \|\theta\|_2^2 \quad (5)$$

and plot the validation losses for the choices of datasets and λ . You will observe the double descent phenomenon for relatively small λ but not the optimal choice of λ .

Problem details: We assume that the data are generated as described in part (c) with $d = 500, \sigma = 0.5$. You are asked to have a Gaussian prior $\theta \sim \mathcal{N}(0, \eta^2 I)$ with $\eta = 1/\sqrt{d}$ on the parameter θ . (The teaching staff indeed generated the ground-truth θ from this prior.)

You are given 13 training datasets of sample sizes $n = 200, 250, \dots, 750$, and 800, and a validation dataset, located at

- `src/bayesianreg/train200.csv`, `train250.csv`, etc.
- `src/bayesianreg/validation.csv`

Let λ_{opt} denote the regularization strength that you derived from part (c). (λ_{opt} is a function of σ and η .)

For each training dataset $(X, \vec{y})$ and each $\lambda \in \{0, 1/32, 1/16, 1/8, 1/4, 1/2, 1, 2, 4\} \times \lambda_{\text{opt}}$, compute the optimizer of equation (5), and evaluate the mean-squared-error of the optimizer on the validation dataset.

Complete the `ridge_regression` method of `src/doubledescent/doubledescent.py` which takes in a training file and a validation file, computes the θ that minimizes training objective under different regularization strengths, and returns a list of validation errors (one for each choice of λ).

Include in your writeup a plot of the validation losses of these models. The x-axis is the size of the training dataset (from 200 to 800); the y-axis is the MSE on the validation dataset. Draw one line for each choice of λ connecting the validation errors across different training dataset sizes. Therefore, the plot should contain 9×13 points and 9 lines connecting them.

You should observe that for those choices of λ with $\lambda < \lambda_{opt}$, the validation error may increase and then decrease as we increase the sample size. However, double descent does not occur for λ_{opt} or any regularization larger than λ_{opt} .

Note: Use the Moore-Penrose pseudo-inverse as implemented in `numpy.linalg.pinv` if your matrix is singular.

Remark: If you would like to know more about double descent, please feel free to check out the partial list of references in the introduction and related work in [?], but knowing the references or papers are unnecessary for solving this homework problem. Roughly speaking, it is mostly caused by the lack of regularization when $n \approx d$. When $n \approx d$, the data matrix is particularly ill-conditioned and stronger and more explicit regularization is needed.

- (e) [5 points] Next, consider the Laplace distribution, whose density is given by

$$f_{\mathcal{L}}(z|\mu, b) = \frac{1}{2b} \exp\left(-\frac{|z - \mu|}{b}\right).$$

As before, consider a linear regression model given by $y = x^T \theta + \epsilon$ where $\epsilon \sim \mathcal{N}(0, \sigma^2)$. Assume a Laplace prior on this model, where each parameter θ_i is marginally independent, and is distributed as $\theta_i \sim \mathcal{L}(0, b)$.

Show that θ_{MAP} in this case is equivalent to the solution of linear regression with L_1 regularization, whose loss is specified as

$$J(\theta) = \|X\theta - \vec{y}\|_2^2 + \gamma \|\theta\|_1$$

Also, what is the value of γ ?

Note: A closed form solution for linear regression problem with L_1 regularization does not exist. To optimize this, we use gradient descent with a random initialization and solve it numerically.

Remark: Linear regression with L_2 regularization is also commonly called *Ridge regression*, and when L_1 regularization is employed, is commonly called *Lasso regression*. These regularizations can be applied to any Generalized Linear models just as above (by replacing $\log p(y|x, \theta)$ with the appropriate family likelihood). Regularization techniques of the above type are also called *weight decay*, and *shrinkage*. The Gaussian and Laplace priors encourage the parameter values to be closer to their mean (*i.e.*, zero), which results in the shrinkage effect.

Remark: Lasso regression (*i.e.*, L_1 regularization) is known to result in sparse parameters, where most of the parameter values are zero, with only some of them non-zero.

$$(a) \theta_{\text{map}} = \arg \max_{\theta} p(\theta|x, y) = \arg \max_{\theta} \frac{p(y|x, \theta) p(\theta|x)}{p(y|x)}.$$

since $p(y|x)$ can't be optimized by θ , and $p(\theta|x) = p(\theta)$

$$\therefore \theta_{\text{map}} = \arg \max_{\theta} p(y|x, \theta) p(\theta).$$

$$(b) p(\theta) = \frac{1}{\sqrt{2\pi}\eta} \exp\left(-\frac{1}{2} \frac{(\theta_i - 0)^2}{\eta^2}\right)$$

$$\theta_{\text{MAP}} = \arg \max_{\theta} p(\theta | x, y)$$

$$= \arg \max_{\theta} p(y | x, \theta) \cdot p(\theta)$$

$$= \arg \min_{\theta} \left[-\log p(y | x, \theta) + \underbrace{\frac{1}{2} \frac{\|\theta\|_2^2}{\eta^2 I}}_{\text{terms w/o } \theta \text{ are dropped.}} \right] \text{ from pdf of } \theta.$$

Let $\lambda = \frac{1}{2\eta^2 I}$ then

$$\theta_{\text{MAP}} = \arg \min_{\theta} (-\log p(y | x, \theta) + \lambda \|\theta\|_2^2).$$

(c) $y \in \mathbb{R}^n$, $x \in \mathbb{R}^{d \times n}$, $\varepsilon \in \mathbb{R}^n$, $\theta \in \mathbb{R}^{d \times 1}$

$$\theta_{\text{MAP}}^{(i)} = \arg \min_{\theta^{(i)}} \left(-\log p(y^{(i)} | x^{(i)}, \theta^{(i)}) + \lambda \|\theta^{(i)}\|_2^2 \right)$$

$$p(y | x, \theta) = p(\theta^T x + \varepsilon | x, \theta) = p(\varepsilon) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2} \frac{\varepsilon^2}{\sigma^2}\right)$$

$$\theta_{\text{MAP}}^{(i)} = \arg \min_{\theta^{(i)}} \left(\frac{1}{2} \frac{\varepsilon^{(i)2}}{\sigma^{(i)2}} + \lambda \|\theta^{(i)}\|_2^2 \right)$$

$$= \arg \min_{\theta^{(i)}} \left(\frac{1}{2} \frac{(y^{(i)} - \theta^{(i)T} x^{(i)})^2}{\sigma^{(i)2}} + \lambda \|\theta^{(i)}\|_2^2 \right)$$

$$= \arg \min_{\theta^{(i)}} \left((y^{(i)} - \theta^{(i)T} x^{(i)})^2 + \lambda' \|\theta^{(i)}\|_2^2 \right)$$

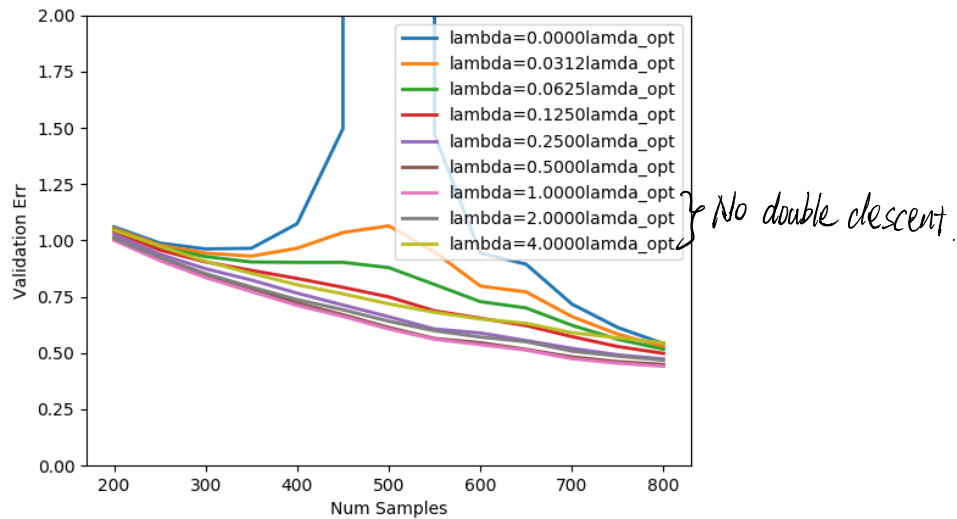
where $\lambda' = \frac{\sigma^{(i)2}}{\eta^2 I}$

This is a regularized least square problem.

We have $Y | X, \theta \sim N(X\theta, \vec{\sigma}^2)$, $\theta \sim N(0, \eta^2 I)$

$$\hat{\theta}_{\text{MAP}} = X^T (X X^T + \frac{\sigma^2}{\eta^2 I})^{-1} Y.$$

(a)



We can see that the last three curves does not seems to have this effect.

$$(b). \quad p(\theta) = \frac{1}{2b} \exp\left(-\frac{|\theta|}{b}\right)$$

$$p(y|x, \theta) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2} \frac{\|y - x\theta\|_2^2}{\sigma^2}\right)$$

$$-\log p(\theta) = \log(2b) + \frac{|\theta|}{b}$$

$$\theta_{\text{MAP}} = \arg \max_{\theta} p(y|x, \theta) \cdot p(\theta)$$

$$= \arg \min_{\theta} \left(-\log p(y|x, \theta) + \frac{\|\theta\|_1}{b} \right)$$

$$= \arg \min_{\theta} \left(\|x\theta - \vec{y}\|_2^2 + \gamma \|\theta\|_1 \right) \quad \text{where } \gamma = \frac{1}{b}.$$